



Semnan University

Journal of Modeling in Engineering

Journal homepage: <https://modelling.semnan.ac.ir/>

ISSN: 2783-2538



Research Article

Storage Location Assignment Problem using Data Mining Approach (Case Study: Mobarakeh Steel Company)

Hossein Raesi Dezaki ^{a,*}, GholamAli Reisi Ardali ^a

^a Department of Industrial and Systems Engineering, Isfahan University of Technology, Isfahan, Iran

PAPER INFO

Paper history:

Received: 2024-02-15

Revised: 2024-09-05

Accepted: 2024-10-19

Keywords:

Inventory management;
Data mining;
Clustering;
GA.

ABSTRACT

One of the influential factors in manufacturing industries to maintain competitiveness and meet customer expectations is optimizing warehouse management, including the allocation of products to storage locations within the warehouse. Proper arrangement can reduce order picking time and increase responsiveness to orders in the warehouse. This research focuses on the topic of inventory management using data-driven tools. In this study, one of the warehouses of Mobarakeh Steel in Isfahan was considered as a case study, and the process of allocating products to shelves was carried out in two stages. In the first stage, products were divided into four clusters using the k-means algorithm. To improve the quality of clustering, two steps were taken: "a pre-algorithm that selects initial k-means points with greater dispersion" and "a genetic algorithm that eliminates local optima in cluster solutions." In the second stage, the location of each cluster in the warehouse was determined based on experts experience. Subsequently, based on a preference criterion, products in each cluster were assigned to different shelves. Several orders were randomly selected, and the total distances of the ordered products were calculated. The results show that the proposed model can reduce the total distances of products in orders in the warehouse by an average of 15 percent.

DOI: <https://doi.org/10.22075/jme.2024.33295.2624>

© 2025 Published by Semnan University Press.

This is an open access article under the CC-BY 4.0 license. (<https://creativecommons.org/licenses/by/4.0/>)

* Corresponding author.

E-mail address: h.raesi@in.iut.ac.ir

How to cite this article:

H. Raesi Dezaki and G. Reisi Ardali, "Storage location assignment problem using data mining approach (Case Study: Mobarakeh Steel Company)," Journal of Modeling in Engineering, 23 82 (2025): 23-35, doi: 10.22075/jme.2024.33295.2624

انبارش کالاها در انبار با رویکرد داده کاوی (مطالعه موردی: شرکت فولاد مبارکه اصفهان)

حسین رئیسی دزکی^{۱*}، غلامعلی رئیسی اردلی^۱

اطلاعات مقاله	چکیده
دریافت مقاله: ۱۴۰۲/۱۱/۲۶ بازنگری مقاله: ۱۴۰۳/۰۶/۱۵ پذیرش مقاله: ۱۴۰۳/۰۷/۲۸	یکی از عوامل اثرگذار در صنایع تولیدی برای حفظ رقابت پذیری و برآورده کردن انتظارات مشتریان مدیریت بهینه انبار شامل تخصیص محصولات به مکان های ذخیره سازی درون انبار می باشد. چیدمان مناسب می تواند، مدت زمان جمع آوری سفارشات را کاهش داده و منجر به افزایش زمان پاسخگویی به سفارشات در انبار گردد. تمرکز این پژوهش بر موضوع انبارش کالا با استفاده از ابزارهای داده محور می باشد. در این پژوهش، یکی از ابزارهای فولاد مبارکه اصفهان به عنوان مطالعه موردی در نظر گرفته شد و فرایند تخصیص کالاها به قفسه ها در دو مرحله صورت پذیرفت. در مرحله اول کالاها با استفاده از الگوریتم k-means به چهار خوشه تقسیم شد و برای افزایش کیفیت خوشه بندی، از دو اقدام: "یک پیش الگوریتم که نقاط اولیه k-means را با پراکندگی بیشتری انتخاب می کند" و "یک الگوریتم ژنتیک که جواب های خوشه بندی را از بهینگی محلی خارج می کند" استفاده گردید. در مرحله دوم با استفاده از تجربه خبرگان، مکان هر خوشه در انبار مشخص شده است. در ادامه بر اساس یک معیار ارجحیت، کالاهای هر خوشه به قفسه های مختلف اختصاص یافته است. چند سفارش به صورت تصادفی انتخاب شده و مجموع فاصله های کالاهای سفارشات محاسبه شد. نتایج نشان می دهد که مدل پیشنهادی می تواند مجموع فاصله های کالاها در سفارشات را در انبار به طور متوسط ۱۵ درصد کاهش دهد.

واژگان کلیدی:

انبارداری،
انبارش کالا در انبار،
خوشه بندی،
الگوریتم ژنتیک.

DOI: <https://doi.org/10.22075/jme.2024.33295.2624>

© 2025 Published by Semnan University Press.

This is an open access article under the CC-BY 4.0 license. (<https://creativecommons.org/licenses/by/4.0/>)

۱- مقدمه

جهانی شدن سریع بازار محصولات و افزایش سفارشی سازی هر محصول، چالش های بسیاری را برای صنایع تولیدی به همراه داشته است. صنعت فولادسازی یکی از صنایع مهم اقتصاد ملی هر کشور به حساب می آید بگونه ای که بدلیل رقابت بالای این صنعت در دنیا با وظیفه دشوار تولید محصولات با کیفیت بالا و توسعه سبز مواجه است. انبارها به عنوان بخش مهمی که مواد اولیه، تولیدات واسطه ای و محصول نهایی کارخانه در آنها نگهداری می شود و همچنین در سیستم مدیریت زنجیره تأمین نقش حیاتی دارند، در بهبود کارایی تولید از اهمیت قابل توجهی برخوردار هستند.

بگونه ای که مسئولیت ورود، ذخیره و مدیریت موجودی تا زمان خروج از انبار را بر عهده دارند [۲]. مدیریت بهینه انبار در صنایع تولیدی برای حفظ رقابت پذیری و برآورده کردن انتظارات مشتریان امری ضروری به حساب می آید. بهینه سازی عملیات های موجود در انبار می تواند منجر به بهبود پایداری، کاهش هزینه ها و افزایش رضایت مشتری شود. یکی از مسائل مهم در مدیریت انبار، مسئله انبارش کالا در انبار (SLAP) می باشد که در آن محصولات به مکان های ذخیره سازی درون یک انبار تخصیص داده می شوند. نحوه تخصیص کالاها به مکان های ذخیره سازی تأثیر مستقیمی بر فرآیندهای جمع آوری سفارشات دارد و

* پست الکترونیک نویسنده مسئول: h.raesi@in.iut.ac.ir

۱. دانشکده مهندسی صنایع، دانشگاه صنعتی اصفهان، اصفهان، ایران

استناد به این مقاله:

دادند. این مدل از یک مدل ریاضی برای خوشه‌بندی استفاده کرده و با استفاده از تکنیک‌های چند معیاره، محصولات را به انبارها تخصیص داده است. نتایج مطالعه بهبود ۴۵٪ در کاهش مسافت جمع‌آوری سفارشات را نشان می‌دهد.

فانتا و همکاران [۷]، یک فرآیند ذخیره‌سازی مبتنی بر کلاس و اختصاص مکان ذخیره‌سازی بر اساس شاخص مکعب به ازای هر سفارش (COI) پیشنهاد می‌دهد. این فرآیند با استفاده از محاسبات پارتو، فضای مورد نیاز و کل فاصله برداشت سفارش را تحلیل می‌کند.

داسیلوا و همکاران [۸] یک مدل تصمیم‌گیری چند معیاره برای رتبه‌بندی محصولات و تخصیص آنها به مکان‌های ذخیره‌سازی در انبار ارائه دادند. روش حل به چند مرحله تقسیم می‌شود، ابتدا بهترین مکان‌ها در انبار شناسایی می‌شوند. سپس، با توجه به شاخص‌های مسئله، کالاهای ناچیره بر اساس روش چند معیاره SMARTER رتبه‌بندی می‌شوند. تخصیص این محصولات بر اساس رتبه‌بندی آنها و ظرفیت ذخیره‌سازی قفسه‌های مختلف است و همچنین، کالاهای چیره بر اساس روش لکسیکوگراف تخصیص داده می‌شوند. نتایج داده‌های فرضی نشان‌دهنده افزایش سرعت در پردازش سفارشات نسبت به روش‌های سنتی موجود در انبار می‌باشد.

لی و همکاران [۹] یک الگوریتم داده‌کاوی توسعه داده‌اند (یک روش اکتشافی مبتنی بر رابطه محصولات که PABH نامگذاری شده است) که رابطه بین هر جفت محصولات را استخراج کرده و سپس با توجه به این داده‌ها (ماتریس جریان بین تسهیلات) و اطلاعات موجود از انبار (ماتریس فاصله)، یک مسئله تخصیص درجه دو تعریف شده است و از الگوریتم ژنتیک برای حل مسئله استفاده شده است. مقایسه عملکرد بین رویکرد جدید و روش طبقه‌بندی سنتی ABC انجام شده است که مطالعه موردی انجام‌شده بر روی یک مرکز توزیع تولیدکننده محصولات مراقبت بهداشتی، بهبود ۷.۱۴ تا ۱۰۴.۴ درصدی را در میانگین زمان انتخاب سفارشات نشان می‌دهد.

پانگ و چان [۱۰] یک الگوریتم مبتنی بر داده‌کاوی را برای تعیین محل ذخیره‌سازی اقلام در یک انبار تصادفی از طریق استخراج، تجزیه و تحلیل وابستگی بین محصولات مختلف در سفارشات مشتری ارائه داده‌اند. هدف این الگوریتم به حداقل رساندن مجموع فواصل سفر برای دسترسی به کالاها

با بهینه‌سازی چیدمان کالاها در انبار، فاصله بین کالا‌های سفارشات کاهش یافته و سرعت پاسخگویی به سفارشات افزایش می‌یابد. رویکردهای متنوعی برای حل مسئله SLAP وجود دارد که از جمله آن می‌توان به مدل‌سازی ریاضی، الگوریتم‌های هیوریستیک و تکنیک‌های تجزیه و تحلیل داده اشاره کرد. یکی از تکنیک‌ها و ابزارهای مؤثر برای توسعه و بهبود مسئله انبارش کالا در انبار استفاده از روش‌های داده‌کاوی می‌باشد. داده‌کاوی فرایندی است که حجم قابل‌توجهی از داده‌ها را به اطلاعات معنی‌دار برای پشتیبانی کارآمد از تصمیم‌گیری تبدیل می‌کند. به عبارت دیگر به مجموعه‌ای از روش‌های قابل اعمال بر پایگاه داده‌های بزرگ و پیچیده به منظور کشف الگوهای پنهان و جالب‌توجه نهفته در بین داده‌ها، داده‌کاوی گویند [۳]. ابزارهای مختلفی جهت داده‌کاوی وجود دارند که تمرکز هر کدام بر پیش‌بینی، طبقه‌بندی و خوشه‌بندی داده‌ها می‌باشد [۴]. با مطالعه مقالات بروز داده‌کاوی با موضوع نحوه چیدمان کالا در انبار و راه‌های بهبود آن که با اهداف مختلفی همچون بهینه‌سازی فضای انبار، پیش‌بینی تقاضا، تامین و توزیع کالا، پیش‌بینی مدیریت موجودی و غیره انجام‌شده است، می‌توان به اهمیت و کاربردی بودن موضوع علم داده در بهبود عملکرد این بخش پی برد [۵]. تمرکز این پژوهش بر موضوع انبارش کالا با استفاده از ابزارهای داده محور می‌باشد و هدف در این پژوهش ارائه مدلی جهت چیدمان کالاها در انبار به منظور کاهش مجموع فاصله بین کالاهای سفارشات می‌باشد.

۲- بررسی ادبیات موضوع

محققان از رویکردها و تکنیک‌های مختلفی برای حل مسئله انبارش کالا در انبار استفاده می‌کنند. از جمله تکنیک‌های استفاده شده برای حل این مسئله می‌توان به برنامه‌ریزی عدد صحیح، الگوریتم‌های فراابتکاری، شبیه‌سازی، تصمیم‌گیری چند معیاره و داده‌کاوی اشاره کرد که هر کدام شامل یک حوزه تحقیقاتی مجزا بوده و چالش‌های خاص خود را دارا می‌باشد. با توجه به پژوهش مروری رنژ و همکاران [۵]، نیاز به تحقیقات بیشتر در مورد توسعه ابزارهای تصمیم‌گیری از طریق استفاده از تکنیک‌های یادگیری ماشین و رویکردهای ترکیبی مبتنی بر داده‌کاوی توصیه شده است.

ییفو و همکاران [۶] یک مدل برای خوشه‌بندی و تخصیص اقلام به منظور کاهش مسافت جمع‌آوری سفارشات ارائه

می باشد. آزمایش های محاسباتی گسترده مبتنی بر داده های انبار قطعات یدکی یک شرکت محصولات کامپوتری و شبکه سازی در هنگ کنگ برای ارزیابی کارایی و کاربرد الگوریتم پیشنهادی انجام شده است. نتایج نشان می دهد که الگوریتم پیشنهادی در صورتی که کالاها ارتباط قوی با یکدیگر داشته باشند، برای بهبود بهره وری عملیات ذخیره سازی قابل اعمال است.

چانگ و همکاران [۱۱] یک مدل جدید را ارائه می دهند که الگوی همبستگی تقاضا را در مسئله SLAP در نظر می گیرد. این مدل به همبستگی مکانی و زمانی تقاضا برای اقلام توجه می کند. همبستگی مکانی به تأثیر اقلام همسایه بر تقاضای یکدیگر اشاره دارد، در حالی که همبستگی زمانی به تأثیر تقاضای گذشته بر تقاضای آینده اشاره دارد. آنها یک مدل عدد صحیح مختلط ارائه و همچنین آزمایش های محاسباتی را مبنی بر ارزیابی عملکرد مدل پیشنهادی انجام داده و آن را با نمونه های عددی مقایسه می کنند.

مایکل و همکاران [۱۲] بیان داشتند که مسئله SLAP در ادبیات موضوع با ماهیت چند معیاره بودن شناخته می شود به همین منظور یک روش ترکیبی مبتنی بر روش های چند معیاره به طور خاص (ELECTRE TRI و TOPSIS) در مقاله حاضر طراحی شده است. همچنین برای در نظر گرفتن عدم قطعیت در فرآیند تصمیم گیری، داده ها تحت یک بازه فاصله ای تعریف شده اند. در این روش حل، ابتدا از ELECTRE TRI برای تخصیص محصولات به قفسه استفاده شده و سپس از TOPSIS برای تعیین محل دقیق ذخیره سازی آنها در هر قفسه به جای استفاده از یک سیاست تخصیص تصادفی استفاده می شود.

لورنس و لهر [۱۳] مدلی توسعه داده اند که ابتدا داده های ورودی شامل لیست های انتخاب سفارش مشتریان، ویژگی های فیزیکی کالاها و ساختار انبار پردازش می شود. داده های ورودی برای یافتن همبستگی بین محصولات تجزیه و تحلیل می شوند. همچنین از روش های استاندارد دسته بندی محصولات (تحلیل ABC، تحلیل XYZ) استفاده می شود. سپس نتایج دسته بندی محصولات با ویژگی های محصولات و مقدار تقاضا آنها با تجزیه و تحلیل آماری مقایسه می شود. در مرحله بعد، بهترین نتایج دسته بندی و مهمترین معیارها بدست آمده و برای آموزش شبکه عصبی مصنوعی (پرسپترون دولایه) استفاده می شود. در انتها شبکه عصبی آموزش دیده، به منظور تخصیص

کالاها مورد استفاده قرار می گیرد. نتایج، کاهش ۱۰ تا ۱۶ درصدی کل زمان برداشت سفارش با استفاده از مدل پیشنهادی را نشان می دهد.

یرلیکایا [۱۴] با روش فازی پرموت تحت معیارهای کیفی همچون: مقدر تقاضا، مقدار سود حاصله کالا و میزان حساسیت مشتری نسبت به کالا رتبه بندی شده اند و با توجه به این رتبه ها به مناسب ترین مکان های ذخیره سازی اختصاص داده می شوند. اثربخشی روش پیشنهادی با یک نمونه کوچک آزمایش شده است.

لورنس و همکاران [۱۵] رویکرد جدیدی برای دسته بندی کالاها و پیش بینی وضعیت آنها ارائه داده اند. ابتدا لیست انتخاب سفارشات مشتریان، تجزیه و تحلیل شده و معیارهای مورد نیاز اعم از: تعداد سفارش های هر کالا در ماه، موجودی کالاها و تفاوت بین تعداد سفارشات این ماه با ماه قبلی استخراج شده است. سپس این اطلاعات به عنوان داده های ورودی به خوشه بندی (k-means) داده شده و سپس با استفاده از ماتریس فاصله، میانگین فواصل برای سفارشات محاسبه و از داده های برجسته شده برای ورودی شبکه عصبی مصنوعی (پرسپترون دو لایه) استفاده شده است. از مدل شبکه عصبی آموزش داده شده برای تخصیص کالاها به انبار استفاده شده است. در ادامه به منظور ارزیابی از داده های واقعی ۳۸۰۰۰ سفارش برای ماه های سپتامبر تا ژانویه استفاده شده و نمودارهای box&plot برای وضعیت موجود انبار و مدل های پیشنهادی ترسیم شده است. نتایج بدست آمده نشان می دهد که استفاده از خوشه بندی می تواند کارایی فرآیند انتخاب سفارش را افزایش دهد.

جو و همکاران [۱۶]، بهبود استراتژی انبارش کالا در انبار را از طریق استفاده از تکنیک های داده محور، از جمله جمع آوری داده ها، تجزیه و تحلیل خوشه ای و تجزیه و تحلیل ارتباط بین کالاها، برای بهبود کارایی انتخاب سفارش بررسی کرده اند.

الگوریتم خوشه بندی k-means برای دسته بندی انواع کالاها در سفارشات استفاده شده است. خوشه بندی بر اساس میزان گردش کالا، ارزش کالا، حجم فروش، ارسال رایگان و پشتیبانی نقدی هنگام تحویل انجام می شود. تعداد خوشه ها با توجه به تکنیک سنتی ABC، ۳ خوشه در نظر گرفته شده، سپس الگوریتم آپروری برای تعیین همبستگی بین کالاها و استخراج قواعد وابستگی بین آنها استفاده

می‌شود. مجموعه‌های دو تایی از اقلام مکرر برای هر خوشه استخراج شده و همچنین به منظور تخصیص کالاها به انبار روشی برای بهبود استراتژی ذخیره سازی مبتنی بر کلاس پیشنهاد شده است (کوتاه‌ترین زمان دسترسی به مکان‌های ذخیره‌سازی که توسط خود مقاله تعریف شده است). سپس به منظور ارزیابی مدل پیشنهادی استراتژی ذخیره سازی بهبود یافته با راهبرد سنتی ABC از طریق شبیه‌سازی مقایسه و نشان داده شده‌است که کارایی پردازش سفارشات با استراتژی بهبود یافته افزایش می‌یابد.

کوچک‌دینیز و سونمز [۱۷] به توسعه روشی برای انتخاب سفارش مبتنی بر خوشه‌بندی فازی (FCM) برای انبارش کالاها در انبار می‌پردازد. هدف از روش پیشنهادی به حداقل رساندن کل مسافت طی شده و کل وزن حمل شده برای هر سفارش در عملیات انبار است. به این منظور دو شاخص، وزن کالاها و شاخص Q (تعریف شده توسط محقق) که هم تعداد سفارش کالاها و هم پراکندگی سفارشات را در طول سال نشان می‌دهد، بیان شده‌اند. سپس مسئله با توجه به این دو شاخص خوشه‌بندی شده و خوشه‌ها رتبه‌بندی می‌شوند. با توجه به رویکردهای مرسوم موجود در مدیریت انبار (تخصیص کالاها از نزدیکترین مکان تا نقاط ورودی/خروجی انبار) و خروجی‌های مرحله قبل تخصیص انجام می‌گیرد. یک دیتاست تجربی از داده‌های یک انبار تهیه شده و ارزیابی صورت گرفته نشان‌دهنده مطلوبیت روش پیشنهادی نسبت به روش‌های سنتی موجود در ادبیات موضوع می‌باشد.

بررسی ادبیات موضوع نشان می‌دهد که در مسئله مذکور دو گام اصلی وجود دارد. ابتدا کالاها با در نظر گرفتن معیارهایی، تقسیم‌بندی شده و در کلاس‌های مختلف قرار می‌گیرند و در مرحله بعد با در نظر گرفتن هدف مسئله، کالاها به قفسه‌های انبار تخصیص داده می‌شوند که در هر مرحله از ابزارهای تصمیم‌گیری مختلفی استفاده شده است. با توجه به بررسی‌های انجام شده خلاءها و نوآوری‌های زیر مدنظر می‌باشد و در این پژوهش تلاش می‌شود که مورد بررسی قرار بگیرد.

(۱) انتخاب ویژگی: در تمامی پژوهش‌های صورت گرفته، استدلالی مبنی بر نحوه انتخاب ویژگی‌های تصمیم‌گیری برای دسته‌بندی کالاها وجود ندارد و در هیچ پژوهشی از نظر خبره استفاده نشده است. در این پژوهش ابتدا ویژگی‌های پراهمیت ادبیات موضوع استخراج می‌شود و

سپس با نظر خبره فرایند انتخاب ویژگی انجام می‌شود. (۲) استفاده از تکنیک‌های هیبریدی: همانطور که بیان شد، گام‌های اصلی پژوهش به ترتیب خوشه‌بندی کالاها و تخصیص کالاها در هر خوشه می‌باشد که در این پژوهش با ارائه یک روش ترکیبی، تلاش می‌شود کیفیت خوشه‌بندی بهبود یابد.

(۳) مطالعه موردی انبار فولاد مبارکه اصفهان: در این پژوهش، یک از انبارهای فولاد مبارکه اصفهان مدنظر می‌باشد و مدل ارائه شده با وضعیت فعلی این انبار مقایسه می‌شود.

۳- مدل ارائه شده

مراحل انجام پژوهش حاضر در چهار بخش کلی شامل انتخاب ویژگی و جمع‌آوری داده‌ها، خوشه‌بندی کالاها، تخصیص کالاها به انبار و ارزیابی مدل ارائه شده است.

➤ گام اول، ویژگی‌های مورد نیاز برای تحلیل داده‌ها انتخاب می‌شوند.

➤ گام دوم، کالاها به گروه‌های مختلف خوشه‌بندی می‌شوند.

➤ گام سوم، با استفاده از تجربه خبرگان، کالاها به مکان‌های ذخیره سازی در انبار تخصیص داده می‌شوند.

۳-۱- ویژگی‌های منتخب

با توجه به پژوهش‌های بررسی شده که در فصل دوم به آنها اشاره شده است، لیستی از ۱۱ ویژگی بر تکرار جمع‌آوری شده که در ادامه مورد بررسی قرار گرفته و منابع مذکور نیز بیان می‌شود.

دوره‌ی موجودی کالاها: نسبت تعداد کالاهای فروخته شده به تعداد کالاهای موجود در انبار را نشان می‌دهد. این ویژگی عملکرد فروش و محبوبیت کالاها را بازتاب می‌دهد [۱۶].

ارزش کالاها: این مقدار میزان سودآوری و ارزش کالاها را نشان می‌دهد [۱۳، ۱۴].

ارائه‌ی ارسال رایگان: یک ویژگی باینری است که نشان می‌دهد آیا فروشنده خدمات تحویل رایگان برای کالاها ارائه می‌دهد یا خیر. این ویژگی رقابت‌پذیری و جذابیت کالاها را نمایان می‌سازد [۱۶].

ضریب Q کالاها: این ویژگی با ضرب تعداد سفارشات که شامل یک کالای خاص هستند در انحراف معیار تاریخ

کالای خاص هستند، اشاره دارد.

درجه تکرار یک کالا: این ویژگی تعداد تکرار یک کالا را در مجموعه اقلام مکرری که حاوی آن کالا هستند، نشان می‌دهد. برای محاسبه این ویژگی، ابتدا با استفاده از الگوریتم آپریوری [۲۰]، مجموعه‌اقلام مکرری که دارای Minsup بزرگتر از ۱٪ هستند استخراج می‌شود. سپس ویژگی مدنظر از شمارش تعداد تکرار هر کالا در کل مجموعه اقلام مکرر بدست می‌آید. در ادامه در جدول شماره ۱-۳ نمادهای بکار رفته در الگوریتم و در شکل ۳-۲ شبه کد الگوریتم درجه تکرار یک کالا قابل مشاهده می‌باشد.

تقاضای کالا: تعداد کالای مصرف شده در ۱۲ ماه گذشته است که نشان‌دهنده میزان نیاز سازمان به یک کالای خاص می‌باشد.

ارزش کالاها: به میزان درجه اهمیت کالاها از نظر تأمین و تدارکات اشاره دارد که براساس معیارهای میزان پیش‌بینی تقاضا در سال آینده، مقدار کمبود در سال‌های گذشته و اهمیت کالا از نظر سختی تأمین مشخص می‌شود. این ویژگی با نظر خبرگان و با استفاده از طیف دسته‌ای نه تایی مشخص می‌شود. در جدول ۲ خلاصه ویژگی‌های مورد استفاده در این پژوهش قابل مشاهده می‌باشد.

جدول ۲- ویژگی‌های مسئله خوشه‌بندی

تعریف	ویژگی
تعداد کالای مصرف شده در ۱۲ ماه گذشته	تقاضای کالا
تعداد کل سفارشات که شامل یک کالای خاص هستند	فراوانی سفارشات کالاها
تعداد تکرار یک کالا را در مجموعه اقلام مکرر	درجه تکرار یک کالا
میزان درجه اهمیت کالاها در تأمین و تدارکات	ارزش محصول

همچنین لازم به ذکر است به منظور تعیین و تخصیص کالاها به انبار از ویژگی‌های حجم هر کالا و موقعیت فیزیکی آن‌ها نیز استفاده شده است که در ادامه توضیح داده می‌شوند.

حجم هر کالا: حجم هر کالا به مقدار فضایی اشاره دارد که هر کالا در انبار اشغال می‌کند. به منظور تعیین این عدد از حداکثر موجودی در یکسال گذشته استفاده شده و با توجه به نظر خبره، حجم هر کالا و ظرفیت هر موقعیت انبار، بدست می‌آید. به عنوان مثال کد کالای ۳۰۰۰۱۱۱۲۰۵۹۵

سفارشات به دست می‌آید. شاخص مدنظر نسبت به فراوانی سفارشات کالاها و هم پراکندگی سفارشات در طول سال حساس است. این ویژگی تقاضا و فصلی بودن کالاها را نشان می‌دهد [۱۷].

وزن کالاها: وزن کلی کالاها را نمایان می‌سازد. این ویژگی نشان‌دهنده سختی و هزینه‌های مرتبط با کنترل و حمل و نقل کالاها می‌باشد [۱۲، ۱۳، ۱۸].

فراوانی سفارشات کالاها: تعداد سفارشات که شامل یک کالای خاص می‌باشد [۱۳، ۱۵].

پخش سفارشات در طول سال: انحراف معیار تاریخ سفارش‌ها می‌باشد. این مقدار فصلی بودن و تغییرپذیری کالاها را نمایان می‌سازد [۱۳، ۱۵].

تقاضای کالا: این ویژگی با توجه به مسئله قابل تعریف است، ولی به صورت کلی مقدار مورد نیاز از یک قلم کالا را تقاضای آن قلم کالا می‌نامند [۱۰، ۱۲، ۱۴، ۱۶، ۱۸].

حساسیت کالاها: درجه آسیب‌پذیری یا شکنندگی کالاها را نشان می‌دهد. این مقدار، کیفیت و دوام کالاها را بازتاب می‌کند [۱۰، ۱۲، ۱۴، ۱۸، ۱۹].

حجم کالاها: فضایی است که توسط کالاها اشغال می‌شود. این ویژگی ظرفیت ذخیره‌سازی و بهره‌برداری از کالاها را نمایان می‌کند [۱۰، ۱۲، ۱۶-۱۹].

موجودی کالاها: این معیار نشان‌دهنده موجودی کالاها در انبار می‌باشد [۱۲، ۱۸-۱۶].

در این پژوهش، با توجه به ادبیات موضوع و لیست ویژگی‌های پر تکرار و مشورت با خبرگان این صنعت ویژگی‌های مناسبی برای تحلیل داده‌ها انتخاب شده‌اند که در ادامه توضیح داده می‌شوند. خصوصیات خبرگان در جدول ۱ قابل مشاهده می‌باشد.

جدول ۱- مشخصات خبرگان

تخصص	سابقه کار
رئیس واحد انبارها و واردات	۱۲
کارشناس واحد انبارها و واردات	۵
سرگروه انبارهای مواد مصرفی	۶
تکنسین کنترل موجودی کل انبارها	۱۳

فراوانی سفارشات کالاها: فراوانی یک کالا در سفارشات، به میزان تکرار و حضور آن کالا در سفارشات مختلف اشاره دارد. این ویژگی به تعداد کل سفارشات که شامل یک

رفته در الگوریتم و در شکل (۲) شبه کد الگوریتم قابل مشاهده می‌باشد.

جدول ۳- نمادهای الگوریتم تعیین نقاط اولیه k-means

نماد	تعریف
Data	مجموعه داده‌ی ورودی
K	تعداد خوشه‌ها
max-iterations	حداکثر تعداد تکرار انجام الگوریتم
best-centers	لیستی از بهترین مراکز خوشه
best-center-distances-sum	بیشترین مجموع فواصل بین مراکز خوشه‌بندی
Distances	لیستی از فواصل بین نقاط داده و مراکز موجود در لیست
first-center	مرکز اولیه که به صورت تصادفی انتخاب می‌شود.
max-distance	بیشترین فاصله نقطه داده از مرکز خوشه
total-distance	مجموع فواصل بین همه جفت نقاط داده
total-center-distance	مجموع فواصل بین همه جفت مراکز موجود در لیست

```

PreAlgorithm: Initialization For K-means Algorithm
1: input data, k, max-iterations
2: best-centers = empty list
3: best-center-distances-sum = 0
4:
5: for iteration from 1 to max-iterations:
6:   centers = empty list
7:   first-center = choose a random data point
8:   centers.append(first-center)
9:
10:  for i from 2 to k:
11:    distances = empty list
12:    for each center in 'centers' list:
13:      | distances.append(maximum distance to any data point * 2)
14:    probabilities = distances / sum of 'distances' list
15:    next-center_index = choose a random index based on 'probabilities'
16:    centers.append(data[next-center_index])
17:  total-center-distance = 0
18:  for i from 1 to k:
19:    for j from i + 1 to k:
20:      | distance = distance between centers[i] and centers[j]
21:      | total-center-distance = total-center-distance + distance
22:
23:  if total-center-distance > best-center-distances-sum:
24:    best-center-distances-sum = total-center-distance
25:    best-centers = centers
26:
27: Output: return best-centers
  
```

شکل ۲- شبه کد الگوریتم تعیین نقاط اولیه k-means

۳-۲-۲- الگوریتم ژنتیک

از آنجایی که الگوریتم k-means یک الگوریتم ابتکاری بوده و با مشکل بهینگی محلی درگیر است و این موجب می‌شود عملکرد خوشه‌بندی تحت تأثیر قرار گیرد، به همین دلیل از الگوریتم ژنتیک به منظور رفع این مشکل استفاده خواهد شد. الگوریتم ژنتیک طراحی شده برای حل مسئله‌ی مطرح شده دارای یک رشته جواب است. کد گذاری این رشته جواب به این صورت می‌باشد که طول رشته جواب به تعداد داده‌های موجود برای خوشه‌بندی می‌باشد که هر ژنوم از رشته نشان‌دهنده شماره خوشه هر داده می‌باشد. برای مثال

که نشان دهنده پیچ شش گوش دین ۷۹۶۸ کلاس ۱۰ می‌باشد دارای حداکثر موجودی ۴۹۵ عدد می‌باشد که با توجه به نظر کارشناس هر موقعیت انبار ظرفیت ۵۰۰ عدد از این قلم کالا را دارا می‌باشد، بنابراین حجم این کالا ۱ واحد می‌باشد. با توجه به توضیحات ذکر شده حجم هر کالا از رابطه (۱) محاسبه می‌شود.

$$(۱) \quad \left[\frac{\text{حداکثر حجم کالا}}{\text{ظرفیت کالا}} \right] = \text{حجم هر کالا}$$

موقعیت فیزیکی: نشان دهنده موقعیت فیزیکی یک کالا در انبار می‌باشد. این موقعیت براساس یک شناسه سه قسمتی که هر قسمت نشان‌دهنده شماره ردیف، شماره ستون و شماره ردیف است اعلام می‌گردد. در شکل (۱) زیر به عنوان مثال یک نمونه ذکر شده است.



شکل ۱- شماره گذاری کالاهای در انبار

۳-۲-۳- خوشه‌بندی کالاهای

در این تحقیق، برای خوشه‌بندی کالاهای از ترکیب دو الگوریتم یعنی k-means بهبودیافته و ژنتیک استفاده شده است. هدف اصلی این ترکیب، ارتقاء دقت و کارایی در فرآیند خوشه‌بندی کالاهای می‌باشد. به منظور خوشه‌بندی ابتدا از الگوریتم k-means استفاده شده و بهترین خروجی‌ها به عنوان ورودی الگوریتم ژنتیک استفاده می‌شود و به عبارت دقیق‌تر پس از مرحله اول از الگوریتم k-means به تعداد مورد نیاز برای تولید جمعیت اولیه در الگوریتم ژنتیک استفاده می‌شود.

۳-۲-۱- الگوریتم k-means

در ابتدا، الگوریتم k-means بهبود یافته به منظور تولید یک مجموعه اولیه از جواب‌ها به کار گرفته می‌شود که نکته مهم در این الگوریتم این است که در انتخاب مراکز اولیه خوشه‌ها تلاش شده بررسی فضای حل با پراگندگی بیشتری در انتخاب مراکز اولیه صورت گیرد. این توازن بر اساس بررسی فضای جستجو و شناخت بهتر از داده‌ها انجام می‌شود و همچنین با تکرار چندین باره این جستجو جواب‌های بهتری پیدا شده و احتمال تصادفی بودن مراکز اولیه کاهش می‌یابد. در ادامه در جدول ۳ نمادهای بکار

۳-۳- تخصیص کالاهای به انبار

پس از انجام عملیات خوشه‌بندی کالاهای، مرحله بعدی تخصیص این کالاهای به انبار است. برای این منظور، ابتدا مراکز هر خوشه، به عنوان نماینده خوشه در نظر گرفته شده و سپس خوشه‌ها تحلیل شده و تلاش می‌شود الویت‌بندی برای خوشه‌ها در نظر گرفته شود. در ادامه با توجه به رابطه (۶) مقدار ارجحیت هر کالا محاسبه می‌شود و به این ترتیب کالاهای هر خوشه به قفسه‌ها تخصیص داده می‌شوند.

$$(۶) \quad \text{درجه تکرار کالا} \times ۰.۷ + \text{تقاضای کالا} \times ۰.۳ = \text{ارجحیت}$$

۴- مطالعه موردی و اجرای الگوریتم‌ها

در شرکت فولاد مبارکه کالاهای در چهار حوزه اصلی تقسیم‌بندی می‌شوند: مواد اولیه و انرژی، تجهیزات و قطعات یدکی، مواد مصرفی، و خدمات و امور پیمانکاری. از میان این دسته‌بندی کالاهای، مواد اولیه، تجهیزات و مواد مصرفی در داخل فولاد مبارکه نگهداری می‌شوند. مساحت فضاهای موجود به همراه فضاهای توسعه کلیه انبارها بالغ بر ۱۵۰۰۰۰ مترمربع است. فرآیند انبارداری شامل دریافت، نگهداری و تحویل کالا به متقاضی است. هنگام ورود، کالا از طریق اسکن بارنامه شناسایی می‌شود و پس از شناسایی، مجوز ورود صادر می‌شود و کالا دریافت می‌شود. سپس اقلام کالاهای بازرسی شده و تحویل انبار داده می‌شود. کالا در محوطه قرنطینه نگهداری می‌شود و پس از تایید کیفیت، وارد انبار می‌شود. انباردار کالا را بررسی کرده و موقعیت کالا در انبار را ثبت می‌کند. در انبار مواد مصرفی، ۴ زیربخش وجود دارد که شامل لوازم اداری، اتصالات سیم جوش، ابزار آلات و پیچ و مهره و لوازم برقی و الکتریکی است. مطالعه موردی این پژوهش بخش پیچ مهره می‌باشد. انبار پیچ و مهره مورد بررسی، محل نگهداری بیش از ۲۰۰۰ نوع محصول مختلف با مقدار موجودی بیش از ۱۰۰۰۰ عدد می‌باشد. این انبار دارای ۳ راهرو، ۶ ردیف و ۹۷ ستون و هر ستون دارای ۳ قفسه می‌باشد و در هر قفسه به صورت میانگین ۹ موقعیت وجود دارد. که در هر موقعیت یک کازیه وجود داشته که محل نگهداری کالا می‌باشد. در این انبار عرض راهروها یک متر در نظر گرفته شده است و هر قفسه دارای طول ۳ متر و عرض ۱ متر می‌باشد. همچنین هر موقعیت محل نگهداری یک کازیه می‌باشد که به طول ۷۰ سانتی متر و عرض ۳۰ سانتی متر است. شکل (۴) نشان‌دهنده شمای کلی انبار می‌باشد.

رشته جواب نمایش داده شده در شکل (۳) یک جواب برای مسئله‌ی مطرح شده با ۱۲ کالا و ۳ خوشه است. این جواب، یک رشته به طول ۱۲ خانه می‌باشد که هر خانه نشان‌دهنده شماره خوشه هر کالا می‌باشد.



شکل ۳- کد گذاری الگوریتم ژنتیک

در این الگوریتم ابتدا با توجه به خروجی‌های الگوریتم kmeans، به تعداد پارامتر population، جواب به عنوان جمعیت اولیه تولید می‌شود و برای هر عضو جمعیت، احتمال انتخاب به صورت رابطه (۲) محاسبه می‌شود که در آن E مقدار تابع خطا که در رابطه (۳) نشان داده شده است و این رابطه مقدار خطا را محاسبه می‌کند که نشان‌دهنده مجموع مربعات فاصله مراکز هر خوشه با اعضای آن خوشه می‌باشد.

$$Selection_p = \frac{E_p}{\sum_{p' \in population} E_{p'}} \quad (۲)$$

$$E = \sum_i \sum_j \|x_i - c_j\|^2 \quad \forall x_i \in C \quad (۳)$$

سپس تعداد nc جواب به عنوان والد برای عملگر ترکیب دو نقطه‌ای و همچنین تعداد nm جواب به عنوان والد برای عملگر جهش از طریق روش چرخ رولت انتخاب می‌شوند که در رابطه (۴) و (۵) قابل مشاهده می‌باشد.

$$nc = p_c \times population \quad (۴)$$

$$nm = p_m \times population \quad (۵)$$

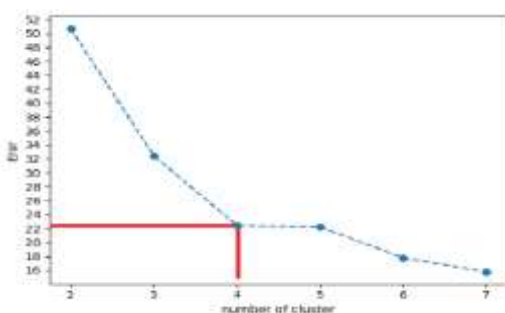
در این روابط پارامترهای pc و pm اعدادی بین صفر و یک بوده و نشان دهنده نرخ جمعیت انتخاب شده برای هر الگوریتم هستند. پس از اجرای عملگرهای ترکیب و جهش برای ایجاد فرزند جدید، بهترین جواب‌ها از میان جواب‌های والد و فرزندان انتخاب شده و به عنوان جمعیت اولیه برای تکرار بعدی الگوریتم استفاده می‌شوند. در این الگوریتم، تابع برازش برابر با رابطه (۳) در نظر گرفته شده و الگوریتم به اندازه‌ی max-iteration مرتبه تکرار می‌شود تا جواب مناسبی به دست آید.

تنظیم می‌شود. در جدول ۵ کلیه پارامترها همراه با توضیحات نشان داده شده است.

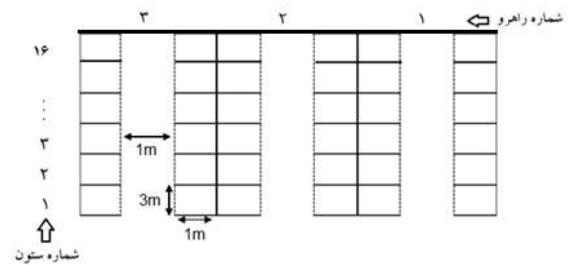
جدول ۵- لیست پارامترهای الگوریتم‌ها

پارامتر	توضیحات
k	تعداد خوشه
Max- GA iteraton	حداکثر تعداد تکرار الگوریتم GA
Population	جمعیت اولیه جواب‌ها
Pc	نرخ انتخاب در عملگر ترکیب
Pm	نرخ جهش

به منظور تنظیم تعداد خوشه‌های مناسب برای الگوریتم k-means از رابطه (۳) استفاده شده است. این رابطه مقدار خطا را محاسبه می‌کند که نشان‌دهنده مجموع مربعات فاصله مراکز هر خوشه با اعضای آن خوشه می‌باشد. با در نظر گرفتن این رابطه پر واضح است که با افزایش تعداد خوشه‌ها (مقدار k)، مقدار خطا کاهش پیدا می‌کند. بنابراین، زمانی که مقدار خطا به یک مقدار ثابت همگرا شد، تعداد مناسب خوشه‌ها مشخص می‌شود. از این رو برای تعیین تعداد خوشه، منحنی تعداد خوشه بر حسب خطا ترسیم می‌شود که این نمودار در شکل (۵) قابل مشاهده است. لازم به ذکر است که نتایج حاصل از این منحنی به ازای هر k سه مرتبه اجرا شده شده است و میانگین هر سه مرتبه به عنوان خروجی (مقدار خطا) در نظر گرفته شده است. و همچنین پارامتر تعداد تکرار برای پیش الگوریتم k-means نیز ۵۰ در نظر گرفته شده است. در نقطه‌ای از محور افقی که شیب منحنی دارای تغییر محسوسی می‌باشد، مناسب‌ترین خوشه‌بندی صورت گرفته است [۲۱] زیرا بعد از این شکستگی، با افزایش تعداد خوشه‌ها، تغییر زیادی در این شاخص رخ نداده است. بنابراین با توجه به توضیحات بالا تعداد خوشه مناسب ۴ در نظر گرفته می‌شود که در شکل (۵) نیز مشخص شده است.



شکل ۵- منحنی تنظیم تعداد خوشه k-means



شکل ۴- شمای کلی انبار

۴-۱- آماده‌سازی پایگاه داده

در این مرحله، ابتدا باید داده‌ها را جمع‌آوری و سازماندهی شوند. همانطور که در فصل قبل بیان شد، ویژگی‌های مناسب انتخاب و مقادیر آنها جمع‌آوری شده و تشکیل یک پایگاه داده را می‌دهند. در جدول ۴ لیست نهایی متغیرها و بازه مقداری آن‌ها قابل مشاهده است.

جدول ۴- لیست نهایی ویژگی‌ها

ویژگی	نوع متغیر	بازه مقداری
تقاضای کالا	عددی	(۰ - ۴۹۷۳۱)
فراوانی سفارشات کالاها	عددی	(۰ - ۱۲۲۰)
درجه تکرار یک کالا	عددی	(۰ - ۲۶۷)
ارزش محصول	دسته‌ای	طیف ۱ تا ۹
حجم هر کالا	عددی	(۰ - ۱۲)

از آنجایی که بازه مقادیر متغیرها متفاوت می‌باشد و برای اینکه الگوریتم‌ها تحت تأثیر بازه مقادیر قرار نگیرد، تمامی ویژگی‌ها به بازه صفر و یک، طبق تابع‌های نرمال‌سازی (۷) نگاشت یافته‌اند. تابع Min-Max یک تبدیل خطی بر روی مقادیر اصلی انجام می‌دهد. این تابع رابطه بین مقادیر داده‌های اصلی را حفظ می‌کند. Min، Max به ترتیب در معادله نشان‌دهنده حداکثر و حداقل مقادیر یک ویژگی می‌باشد. MinNew، MaxNew به ترتیب حداقل و حداکثر بازه جدید می‌باشد.

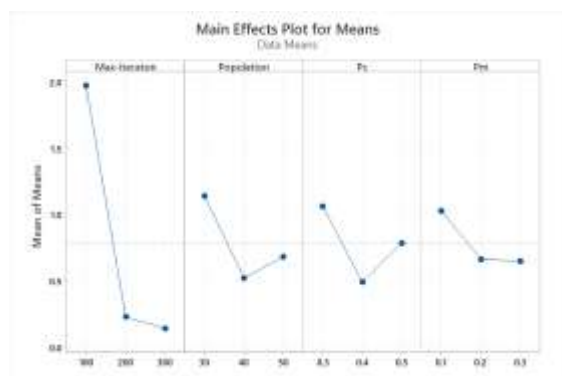
$$newValue =$$

$$\left(\frac{oldValue - Min}{Max - Min} \right) * (MaxNew - MinNew) + MinNew \quad (7)$$

۴-۲- تنظیم و اجرای الگوریتم‌ها

در این پژوهش برای تنظیم پارامتر الگوریتم K-means از شاخص مجموع مربعات خطا [۲۱] و به منظور تنظیم پارامترهای الگوریتم ژنتیک از روش تاگوچی استفاده می‌شود. در تنظیم پارامتر الگوریتم‌های به کار رفته، هر الگوریتم به طور جداگانه بررسی شده و پارامترهای آن

بهرتر است" قرار داده شده است. بنابراین مقادیر مناسب برای پارامترها به گونه‌ای انتخاب می‌شوند که از بین سطوح تعریف شده برای هر پارامتر، نمودار میانگین کمترین مقدار را داشته باشد. در شکل (۵) نمودار میانگین‌های مقادیر به دست آمده برای فاکتورهای مختلف در سطوح مختلف قابل مشاهده است.



شکل ۶- نمودار میانگین آزمایش تاگوچی برای الگوریتم ژنتیک

بر اساس نتایج روش تاگوچی و با توجه به توضیحات ارائه شده، مقادیر مناسب برای چهار پارامتر الگوریتم ژنتیک در جدول ۸ آورده شده است.

جدول ۸- مقادیر تنظیم شده برای الگوریتم ژنتیک

سطح	پارامتر	سطح	پارامتر
۰.۴	Pc	۳۰۰	Max-iteraton GA
۰.۳	Pm	۴۰	Population

بر اساس نتایج حاصل از روش تاگوچی و مقادیر بدست آمده برای پارامترهای الگوریتم ژنتیک، مقدار تابع هدف به عدد ۲۰.۷۸ رسید. این به معنی بهبود حدود ۷.۷ درصد در تابع هدف است. در جدول ۹ تعداد اعضای هر خوشه و مراکز هر خوشه نشان داده شده‌اند.

جدول ۹- نتایج حاصل از خوشه‌بندی برای الگوریتم ژنتیک

مرکز خوشه				تعداد اعضای خوشه	شماره خوشه
ارزش محصول	درجه تکرار یک کالا	فراوانی سفارشات کالاها	تقاضای کالا		
G-H-I	۲.۳	۲۰	۱۳۱۵	۵۱۶	۱
F-G	۰.۴	۱۰.۸	۷۴۵	۶۶۸	۲
C-D-E	۰.۲	۸.۲	۵۳۷	۵۴۹	۳
A-B	۰.۰۴	۴	۲۷۶	۳۹۱	۴

بنابراین با توجه به شکل (۵) تعداد خوشه مناسب ۴ و مقدار خطا برابر با ۲۲.۴ می‌باشد. همچنین در ادامه در جدول ۶ نتایج خوشه‌بندی که شامل تعداد اعضای هر خوشه و مراکز هر خوشه است قابل مشاهده می‌باشد.

جدول ۶- نتایج حاصل از خوشه‌بندی برای الگوریتم k-means

شماره خوشه	تعداد اعضای خوشه	مرکز خوشه		
		تقاضای کالا	فراوانی سفارشات کالاها	درجه تکرار یک کالا
۱	۴۹۷	۱۰۸۴	۱۵.۹	۱.۴
۲	۶۸۷	۹۱۴	۱۳.۸	۱.۱۲
۳	۵۴۹	۵۳۷	۸.۲	۰.۲
۴	۳۹۱	۲۷۶	۴	۰.۰۴

در تنظیم پارامترهای الگوریتم ژنتیک، چهار فاکتور تأثیرگذار وجود دارد که برای هر فاکتور سه سطح در نظر گرفته می‌شود. سطوح مشخص شده برای هر فاکتور در جدول ۷ قابل مشاهده است.

جدول ۷- لیست پارامترهای الگوریتم‌ها

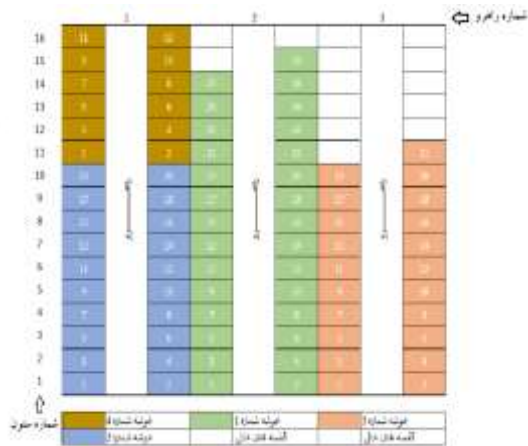
سطح	پارامتر	سطح	پارامتر
۰.۳	Pc	۱۰۰	Max-iteraton GA
۰.۴		۲۰۰	
۰.۵		۳۰۰	
۰.۱	Pm	۳۰	Population
۰.۲		۴۰	
۰.۳		۵۰	

از آنجایی که برای تنظیم پارامترهای الگوریتم ژنتیک توصیه به نرمالسازی نتایج عددی حاصل از جواب‌ها شده است. بنابراین از رابطه (۸) به منظور نرمالسازی استفاده شده است. در این رابطه مقدار جواب در هر آزمایش $Goal_i$ و $NewGoal_i$ مقدار نرمال شده آن جواب را نشان می‌دهد.

$$NewGoal_i = \left(\frac{Goal_i - \min(Goal_1, \dots, Goal_i)}{\min(Goal_1, \dots, Goal_i)} \right) * 100 \quad (8)$$

از آنجایی که تابع هدف تعریف شده برای مسئله مورد نظر کمینه‌سازی است و با توجه به تعریف رابطه‌ی ... مقادیر پارامترها باید به گونه‌ای انتخاب شوند که اختلاف مقدار به دست آمده در هر آزمایش از کمترین مقدار به دست آمده در آزمایش‌های انجام شده برای یک نمونه کمینه گردد. همچنین آزمایش‌های روش تاگوچی بر روی گزینه "کمتر،

موقعیت‌های داخل یک قفسه نیز با یکدیگر یکسان بوده و فرایند تخصیص در حد شماره ستون انجام می‌شود. به عنوان مثال کالاهای ستون اول خوشه شماره ۳ در شکل (۸) قابل مشاهده می‌باشد.



شکل ۷- نقشه تخصیص کالاها به انبار

کد کالا	شرح کالا	حجم کالا	میزان ارجحیت
300011320294	بیج شش گوش یونی ۱۶*۷۰ قفسه/۸ کلاس/۵۷۳۳۹	3	8738.4
300011320292	بیج شش گوش یونی ۱۶*۶۰ قفسه/۸ کلاس/۵۷۳۳۹ (امشابه دین ۹۳۳)	2	3391.9
300013120747	مهدوش شش گوش فلی دین ۱۶*۹۸ کلاس/۵ قفسه/۱۰ کلاس/۵۷۳۳۹ (بسته بندی در بسته های نایلونی 100 عددی)	2	3162.3
300011110722	بیج شش گوش یونی ۱۶*۷۰ قفسه/۸ کلاس/۵۷۳۳۹ (امشابه دین ۹۳۳)	1	3007.4
300011320291	بیج شش گوش یونی ۱۶*۷۰ قفسه/۸ کلاس/۵۷۳۳۹ (امشابه دین ۹۳۳)	2	2350.2
300011110225	بیج شش گوش یونی ۱۶*۷۰ قفسه/۸ کلاس/۵۷۳۳۹ (امشابه دین ۹۳۳)	1	2131.8
300011320298	بیج شش گوش یونی ۱۶*۷۰ قفسه/۸ کلاس/۵۷۳۳۹ (امشابه دین ۹۳۳)	1	1783.4
300013140410	مهدوش شش گوش یونی ۱۶*۷۰ قفسه/۸ کلاس/۵۷۳۳۹ (امشابه دین ۹۳۳)	2	1759.4
300011320371	بیج شش گوش یونی ۱۶*۷۰ قفسه/۸ کلاس/۵۷۳۳۹ (امشابه دین ۹۳۳)	1	1537.4
300011110170	بیج شش گوش یونی ۱۶*۷۰ قفسه/۸ کلاس/۵۷۳۳۹ (امشابه دین ۹۳۳)	1	1523.9
300011110224	بیج شش گوش یونی ۱۶*۷۰ قفسه/۸ کلاس/۵۷۳۳۹ (امشابه دین ۹۳۳)	1	1368.7
300011320369	بیج شش گوش یونی ۱۶*۷۰ قفسه/۸ کلاس/۵۷۳۳۹ (امشابه دین ۹۳۳)	1	1189.4
300012910025	مهدوش شش گوش استیل اسل ۱۶*۷۰ قفسه/۸ کلاس/۵۷۳۳۹ (بسته بندی در بسته های نایلونی 100 عددی)	2	1145.8
300011110168	بیج شش گوش یونی ۱۶*۷۰ قفسه/۸ کلاس/۵۷۳۳۹ (امشابه دین ۹۳۳)	2	1133.9
300011110147	بیج شش گوش یونی ۱۶*۷۰ قفسه/۸ کلاس/۵۷۳۳۹ (امشابه دین ۹۳۳)	1	1131.7
300011320299	بیج شش گوش یونی ۱۶*۷۰ قفسه/۸ کلاس/۵۷۳۳۹ (امشابه دین ۹۳۳)	1	1078.6
300011110245	بیج شش گوش یونی ۱۶*۷۰ قفسه/۸ کلاس/۵۷۳۳۹ (امشابه دین ۹۳۳)	1	1068.7
30001320365	بیج شش گوش یونی ۱۶*۷۰ قفسه/۸ کلاس/۵۷۳۳۹ (امشابه دین ۹۳۳)	1	1042.9
300012130605	بیج سرخونه ای یکجوش خور دین ۱۶*۷۰ قفسه/۸ کلاس/۵۷۳۳۹ (امشابه دین ۹۳۳)	1	965.8

شکل ۸- کالاهای ستون اول خوشه شماره ۳

جدول ۱۱- ارزیابی و مقایسه مدل پیشنهادی با وضعیت فعلی انبار

شماره سفارش	تعداد اقلام	مجموع فواصل کالاها در سفارشات	
		مدل پژوهش	وضعیت فعلی
۱	۶	۵۶	۶۲
۲	۹	۳۱	۴۳
۳	۱۹	۴۶	۴۰
۴	۷	۳۷	۴۳
۵	۱۴	۶۹	۷۵
۶	۱۲	۶۷	۸۲
۷	۸	۴۳	۵۵
۸	۳	۱۲	۳۳
۹	۱۶	۸۵	۷۴
۱۰	۵	۲۲	۳۴
مجموع		۴۶۸	۵۴۱

۴-۳- تخصیص کالاها به قفسه‌ها

با توجه به نتایج خوشه‌بندی و مشخص شدن شماره خوشه هر یک از کالاها و با توجه به ویژگی حجم هر کالا که در بخش ... معرفی شد و نشان‌دهنده تعداد ظرفیت اشغال شده توسط هر کد کالا می‌باشد، می‌توان با استفاده از رابطه (۸) تعیین کرد هر خوشه چه تعداد از قفسه‌ها را اشغال می‌کند که در جدول ۱۰ نشان داده شده است. در این رابطه F_i نشان‌دهنده تعداد قفسه مورد نیاز برای هر خوشه i می‌باشد.

$$F_i = \sum_{j=1}^n \text{حجم هر کالا} \quad (8)$$

جدول ۱۰- تعداد قفسه‌های مورد نیاز

شماره خوشه	تعداد قفسه مورد نیاز
۱	۷۶۰
۲	۵۳۲
۳	۵۴۳
۴	۳۱۲

با توجه به توضیحات بخش ... که شمای انبار مورد نظر را شرح داده است، پر واضح است که مجموعاً این انبار دارای ۲۶۱۹ موقعیت فعال است. همچنین بر اساس نتایج خوشه بندی واضح است که خوشه شماره ۴ که از اهمیت کمتری برخوردار بوده و میزان تقاضای کمتری داشته در گوشه انبار واقع در انتهای راهروی شماره ۱ قرار گرفته زیرا کالاها این خوشه نسبت به دیگر خوشه‌ها از درخواست کمتری برخوردار هستند و خوشه شماره ۱ که از بیشترین میزان تقاضا برخوردار بوده در راهرو شماره ۲ قرار گرفته است تا راحت‌ترین دسترسی به کالاها این خوشه وجود داشته باشد. و دو خوشه دیگر در دو طرف انبار در ابتدای راهروهای شماره ۱ و ۳ قرار گرفته‌اند. در ادامه نقشه تخصیص کالاها به انبار در شکل (۷) قابل مشاهده است. همچنین به منظور تخصیص کالاها به خوشه به قفسه‌ها ابتدا رابطه ... استفاده می‌شود و سپس به ترتیب امتیاز بدست آمده کالاها اولویت‌بندی شده و به نزدیک‌ترین قفسه از ابتدای هر راهرو تخصیص داده می‌شوند. از آنجایی فرض این پژوهش بر دو بعد طول و عرض در نظر گرفته شده است و بعد ارتفاع در نظر گرفته نشده است، بنابراین قفسه‌های موجود در یک ستون از نظر موقعیت مکانی هیچ برتری نسبت به یکدیگر ندارند. و همچنین فرض شده است که

۴-۴- ارزیابی مدل ارائه شده

به منظور ارزیابی مدل ارائه شده ۱۰ سفارش جدید به صورت تصادفی انتخاب شده و مجموع فاصله بین اقلام این سفارشات برای هر دو حالت چیدمان محاسبه شده و همانطور که در جدول ۱۱ قابل مشاهده می باشد به صورت میانگین ۱۵٪ مجموع فواصل کالاها در سفارشات کاهش یافته است.

۵- نتیجه گیری و پیشنهادات

هدف از انجام این پژوهش، ارائه مدلی جهت حل مسئله انبارش کالا در انبار می باشد. که در این مسئله به ارائه سیاستی جهت چیدمان کالاها در انبار با هدف کاهش دادن مجموع فاصله بین کالاها در سفارشات می پردازد. از بررسی های انجام شده در ادبیات موضوع و مقاله مروری رنژ و همکاران [۵] این نتیجه حاصل شد که از ابزارهای داده کاوی جهت حل این مسئله استفاده شود. به همین جهت مدلی ارائه شد که در ابتدا ویژگی های پر اهمیت از ادبیات موضوع جمع آوری شده و سپس با نظر خبره فرایند انتخاب ویژگی انجام شد و در ادامه کالاها خوشه بندی شده و به قفسه ها تخصیص داده شدند. در پایان از یکی از انبارهای فولاد مبارکه به عنوان مطالعه موردی معرفی شد و مدل پیشنهادی با وضعیت فعلی این انبار مورد مقایسه قرار گرفت و نتایج حاکی از کاهش ۱۵

درصدی مجموع فواصل بین کالاها در سفارشات می باشد. برای تحقیقات آینده می توان استفاده از شبکه های عصبی جهت انتخاب ویژگی و همچنین از منطق فازی به منظور مواجهه با عدم قطعیت در تقاضا استفاده کرد.

تقدیر و تشکر

بدینوسیله نویسندگان مقاله، مراتب سپاس و قدردانی خود را از پرسنل واحد انبارها و واردات شرکت فولاد مبارکه، به خاطر مشارکت در گردآوری داده ها، اعلام می دارند.

تعارض منافع

نویسنده اعلام می کند که در مورد انتشار این مقاله تعارض منافع وجود ندارد.

تأییدیه اخلاقی

نویسندگان متعهد میشوند که مطالب این مقاله را در هیچ مجله دیگری به چاپ نرسانده اند.

مشارکت های نویسندگان

حسین رئیسی دزکی: روششناسی، آماده سازی داده ها، انجام مدل سازی، تحلیل نتایج، نگارش مقاله

غلامعلی رئیسی اردلی: راهنمایی علمی، ارزیابی و تفسیر نتایج، اعتبارسنجی نتایج

منابع مالی

ندارد.

مراجع

- [1] Fukuyama, Hirofumi, Hui-hui Liu, Yao-yao Song, and Guo-liang Yang. "Measuring the capacity utilization of the 48 largest iron and steel enterprises in China." *European Journal of Operational Research* 288, no. 2 (2021): 648-665.
- [2] Diabat, Ali, Ehsan Dehghani, and Armin Jabbarzadeh. "Incorporating location and inventory decisions into a supply chain design problem with uncertain demands and lead times." *Journal of Manufacturing Systems* 43 (2017): 139-149.
- [3] Diwani, Salim A, and Zaipuna O Yonah. "A novel holistic disease prediction tool using best fit data mining techniques." *International Journal of Computing and Digital Systems* 6, no. 02 (2017): 63-72.
- [4] Hand, David J. "Principles of data mining." *Drug safety* 30 (2007): 621-622.
- [5] Reyes, J, E Solano-Charris, and J Montoya-Torres. "The storage location assignment problem: A literature review." *International Journal of Industrial Engineering Computations* 10, no. 2 (2019): 199-224.
- [6] Chuang, Yi-Fei, Hsu-Tung Lee, and Yi-Chuan Lai. "Item-associated cluster assignment model on storage allocation problems." *Computers & Industrial Engineering* 63, no. 4 (2012): 1171-1177.
- [7] Fontana, Marcele Elisa, and Cristiano Alexandre Virgínio Cavalcante. "Using the efficient frontier to obtain the best solution for the storage location assignment problem." *Mathematical Problems in Engineering* 2014.
- [8] da Silva, Denilson Dimas, Natália Veloso Caldas de Vasconcelos, and Cristiano Alexandre Virgínio Cavalcante. "Multicriteria decision model to support the assignment of storage location of products in a warehouse." *Mathematical Problems in Engineering* 2015.

- [9] Li, Jiayi, Mohsen Moghaddam, and Shimon Y Nof. "Dynamic storage assignment with product affinity and ABC classification—a case study." *The International Journal of Advanced Manufacturing Technology* 84, no. 9 (2016): 2179-2194.
- [10] Pang, King-Wah, and Hau-Ling Chan. "Data mining-based algorithm for storage location assignment in a randomised warehouse." *International Journal of Production Research* 55, no. 14 (2017): 4035-4052.
- [11] Zhang, Ren-Qian, Meng Wang, and Xing Pan. "New model of the storage location assignment problem considering demand correlation pattern." *Computers & Industrial Engineering* 129 (2019/03/01/ 2019): 210-219.
- [12] Micale, Rosa, Concetta Manuela La Fata, and Giada La Scalia. "A combined interval-valued ELECTRE TRI and TOPSIS approach for solving the storage location assignment problem." *Computers & Industrial Engineering* 135 (2019): 199-210.
- [13] Lorenc, Augustyn, and Tone Lerher. "PickupSimulo—Prototype of Intelligent Software to Support Warehouse Managers Decisions for Product Allocation Problem." *Applied Sciences* 10, no. 23 (2020): 8683.
- [14] Yerlikaya, Mehmet Akif. "Storage Location Assignment with Fuzzy PROMETHEE Method in Warehouse Systems with Uncertain Demand." *Journal of the Institute of Electronics and Computer* 2, no. 1 (2020): 142-150.
- [15] Lorenc, Augustyn, Małgorzata Kuźnar, and Tone Lerher. "Solving product allocation problem (PAP) by using ANN and clustering." *FME Transactions* 49, no. 1 (2021): 206-213.
- [16] Zhou, Li, Lili Sun, Zhaochan Li, Weipeng Li, Ning Cao, and Russell Higgs. "Study on a storage location strategy based on clustering and association algorithms." *Soft Computing* 24, no. 8 (2020): 5486-5499.
- [17] Küçükdeniz, Tarık, and Özlen Erkal Sönmez. "Integrated Warehouse Layout Planning with Fuzzy C-Means Clustering." Paper Presented at the International Conference on Intelligent and Fuzzy Systems, 2022.
- [18] da Silva, Denilson Dimas, Natália Veloso Caldas de Vasconcelos, and Cristiano Alexandre Virginio Cavalcante. "Multicriteria Decision Model to Support the Assignment of Storage Location of Products in a Warehouse." *Mathematical Problems in Engineering* 2015.
- [19] Fontana, Marcela Elisa, and Vilmar Santos Nepomuceno. "Multi-criteria approach for products classification and their storage location assignment." *The International Journal of Advanced Manufacturing Technology* 88, no. 9 (2017): 3205-3216.
- [20] Ye, Yanbin, and Chia-Chu Chiang. "A parallel apriori algorithm for frequent itemsets mining." Paper presented at the Fourth International Conference on Software Engineering Research, Management and Applications (SERA'06), 2006.
- [21] Steinley, Douglas, and Michael J Brusco. "Choosing the number of clusters in K-means clustering." *Psychological Methods* 16, no. 3 (2011): 285-295.